

Context Collapse and the Four Philosophers

Wittgenstein, Lewis, Dennett, Nagel in the Age of AI Chat

Michael Stoyanovich

Artifact Classification: Practitioner Essay

Version 1.4.3 | July 2026

Scope and Limitations

This paper is intended for informational and educational purposes only. The views and analyses presented - particularly those related to ethics, policy, and AI system design - reflect the author's interpretations and do not constitute legal, regulatory, or professional advice. Readers are encouraged to critically assess the content and consult appropriate experts or authorities before applying any concepts discussed herein. The author assumes no liability for any decisions or actions taken on the basis of this work.

Purpose

This practitioner essay examines the problem of context collapse in AI chat systems through the complementary perspectives of Wittgenstein, Lewis, Dennett, and Nagel. It introduces ontological collapse as a practical diagnostic for AI design, governance, and policy, offering conceptual guardrails that help preserve clear distinctions between tools, interlocutors, responsibility, and subjective experience.

Intended Audience

This essay is intended for organizational leaders, technology and governance professionals, policymakers, educators, designers, consultants, and critically engaged readers interested in how generative AI can affect context, interpretation, meaning, and human judgment in mediated communication.

Reading Guide

This essay examines context collapse through several complementary philosophical and practical perspectives, including the Four Philosophers Framework™. It may be read sequentially as an argument about how generative AI can compress, obscure, or reconstruct context, or selectively through the individual conceptual lenses and applied examples. Readers seeking the broader philosophical foundations of the Four Philosophers Framework may consult the related works identified at the end.

Table of Contents

Context Collapse and the Four Philosophers.....	1
Scope and Limitations.....	1
Purpose.....	1
Intended Audience.....	1
Reading Guide.....	1
Abstract.....	3
Introduction.....	3
Executive summary - four guardrail questions for practice.....	4
1. From Audience Collapse to Ontological Collapse.....	5
2. Wittgenstein: Colliding Language-Games in Chat.....	6
3. Lewis: Sham Scorekeeping and Asymmetric Recognition.....	7
4. Dennett: An Engineered Intentional Stance.....	8
5. Nagel: Experience vs. Its Simulation.....	9
6. Productive Tensions: Four Lines of Resistance.....	11
7. Implications for Design, Policy, and Practice.....	11
Checklist - applying the Four Philosophers Framework.....	12
Ethics, Disclosure, and Acknowledgements.....	13
Ethical Considerations.....	13
Use of AI Tools.....	14
Acknowledgements.....	14
Disclosure Statement.....	14
License & Attribution.....	14
Trademark Notice.....	14
How to Cite This Essay.....	14
Related Works.....	14
Version History and Document Status.....	15

Abstract

This paper applies a Four Philosophers framework - Wittgenstein, Lewis, Dennett, and Nagel - to the problem of context collapse in AI chat systems. Whereas classic accounts focus on audience collapse in social media, I argue that large language model interfaces induce a deeper *ontological collapse* between tools, interlocutors, advisors, and experimental subjects. Each philosopher provides a distinct diagnostic: Wittgenstein on colliding language-games, Lewis on sham scorekeeping and misplaced responsibility, Dennett on the engineered intentional stance and competence without comprehension, and Nagel on the boundary between simulated empathy and genuine subjectivity. Together, these lines of analysis yield four guardrail questions for designers, policymakers, and professional users, aimed at resisting downward leveling of human interlocutors while preserving the practical benefits of AI assistance.

Introduction

Social media scholars coined context collapse - see, e.g., boyd and Marwick on social media audiences - to name what happens when messages meant for one audience are suddenly exposed to many: family, colleagues, strangers, all folded into a single, flattened “public.” With large language model (LLM) chat systems, something analogous happens - but now the collapse is not just across audiences, but across roles.

A chat window makes very different practices look the same. In one and the same interface, a user can:

- confide like they would to a friend,
- query like they would a search engine,
- test like they would a research subject,
- and command like they would a tool.

By ontological collapse I mean this: roles that are ordinarily kept distinct - friend, tool, advisor, subject of experiment, bureaucratic system - are compressed into a single interactive object that is treated “as if” it were all of them at once. That is different from:

- epistemic collapse (confusion about what is known or how well), or
- normative collapse (confusion about what is allowed or required).

Ontological collapse is a flattening of what the thing on the other side is taken to be.

The “other side” of the exchange, at present, is not a person, not a subject of experience, and not a member of any normative community. It is a statistical model running on infrastructure. Yet in many deployed systems, the interface is designed to feel like a conversation with “someone”: first-person pronouns, persistent “personality,” simulated empathy. Degrees and motives vary across providers, but anthropomorphic framing generally increases engagement, lowers friction, and keeps attention on the persona rather than the underlying stack.

Users are not passive in this. There are reasons people are ready to lean into the illusion:

- **Loneliness and availability:** a system that is always on, always responsive, and never

visibly bored or impatient is an obvious candidate to receive confessions and anxieties.

- **Convenience and cognitive ease:** speaking “as if” to a person is often less effortful than learning a formal query language or remembering a tool’s exact capabilities.
- **Existing social heuristics:** humans already name ships, yell at laptops, and thank thermostats; chat interfaces simply provide a more persuasive canvas for that tendency.

Context collapse in this setting is therefore overdetermined: commercial incentives push toward anthropomorphic design, and human tendencies pull toward anthropomorphic use.

We can distinguish four interlocking kinds of collapse that show up across social media and AI chat:

Type of collapse	What collapses	Social media example	AI chat example
Audience collapse	Distinct audiences	A post meant for friends goes viral at work	Private-tone prompt screenshotted and shared with a manager
Ontological collapse	What the “other side” is taken to be	Treating a brand account as a friend	Treating an LLM as friend, expert, tool, and subject at once
Epistemic collapse	Confidence and uncertainty distinctions	Equating rumors and reporting in one feed	Treating model output and vetted reference material as equivalent
Normative collapse	Rules and expectations for behavior	Blurred norms across family / work / political circles	Confusing informal “chat” with formal advice or binding decisions

This essay focuses on ontological collapse in AI chat, while keeping the other three in view as background pressure.

This essay is a short conceptual bridge aimed at philosophically literate but non-specialist readers in AI ethics, policy, and design. It uses the Four Philosophers Framework - Wittgenstein (meaning and life), Lewis (norms and recognition), Dennett (competence and comprehension), and Nagel (experience and subjectivity) - as a diagnostic toolkit for ontological collapse in AI chat. The goal is applied: to surface a set of guardrail questions that designers, regulators, and thoughtful users can use to resist that collapse without forfeiting the practical value of these tools.

Executive summary - four guardrail questions for practice

For design, policy, and professional audiences, the practical output of this framework can be expressed as four questions:

1. Wittgenstein - Which language-game are we in?

Are we treating a tool-use interaction as if it were intimate conversation?

2. Lewis - Who is really on the scoreboard?

Are responsibilities being tacitly assigned to a model that cannot bear them?

3. Dennett - Are we mistaking competence for comprehension?

Is a predictive system being treated as if it understood and endorsed its outputs?

4. Nagel - Am I conflating simulated empathy with real subjectivity?

Are human interlocutors being implicitly devalued relative to frictionless simulation?

The four figures do not form a single doctrine; there are real tensions between them, particularly between Dennett's functionalism and Nagel's emphasis on subjective experience. But those tensions generate productive friction: each philosopher highlights a different distinction that the chat interface tends to blur.

1. From Audience Collapse to Ontological Collapse

Classic context collapse is about audiences. A post intended for one circle quietly becomes visible to many; dramaturgical "backstage" and "frontstage" blur; self-presentation becomes a single performance addressed to heterogeneous others.

With LLM-based chatbots, something structurally similar happens - but along a different axis.

When interacting with a chatbot through a messaging-style interface, at least four contexts are silently overlaid:

1. Human-human conversation (empathy, memory, mutual accountability).
2. Search / database query (impersonal tool use).
3. Therapy-like disclosure (expectations of care and discretion).
4. Technical experiment (probing a system with no social obligations).

In one chat log, these modes can follow each other within minutes: asking for a code snippet, then confessing anxiety, then exploring policy scenarios, then benchmarking capabilities. Context collapse here is not only about audiences. It is about what the "other side" is taken to be at any given moment.

The interface nudges users to treat a large-scale statistical model as if it were a conversational partner, a colleague, a helper, sometimes even a confidant. The boundary between person and tool, and between conscious and non-conscious, is softened. People slide back and forth between "talking to it as a thing" and "talking to it as if it were someone," often without noticing the shift.

The Four Philosophers Framework can be read as four ways of recovering structure in that flattened space. If Wittgenstein helps recover the background form of life, Lewis attends to the norms of that life, Dennett analyzes our predictive stances toward systems, and Nagel returns us to the question of who, if anyone, is a subject of experience at all.

The next four sections follow that sequence: from everyday practices and language-games

(Wittgenstein), to norms and commitments (Lewis), to predictive stances (Dennett), and finally to the boundary between experience and its simulation (Nagel).

2. Wittgenstein: Colliding Language-Games in Chat

If context collapse hides who is listening, Wittgenstein helps clarify what game we think we are playing.

Late Wittgenstein ties meaning to use within specific language-games, anchored in particular forms of life. To ask what a word means is to ask how it functions in a shared practice. A promise, a joke, a command, and a diagnosis can all use similar words, but they belong to different games with different stakes.

In LLM chat, users routinely:

- address the system with the grammar of interpersonal dialogue (“Do you understand?”, “Thanks, that helps,” “That must be confusing”),
- while actually engaging in the practice of operating a tool (querying a model, steering an optimization process, triggering latent patterns in a trained network).

The interface encourages a collision of language-games:

- On the surface, users behave as though they are in the game “asking another person for help.”
- In reality, they are in the game “issuing prompts to an engineered system,” with very different criteria for understanding, memory, or sincerity.

Meaning depends on life: the same string of text - “I’m sorry that happened” - is a different kind of move when it comes from a human embedded in shared forms of life than when it is produced by a system sampling from a probability distribution. For the human, that utterance can express genuine regret, carry forward a history, and commit the speaker to future concern. For the model, it is a context-appropriate token sequence chosen because similar strings frequently follow similar prompts in its training data.

Context collapse, in Wittgensteinian terms, is the erasure of that difference. The chat interface hides the underlying form of life (server racks, training runs, policy layers) and foregrounds a familiar conversational surface. What is visible is “chat,” not “remote procedure call into a layered technical system.”

Consider, for example, a patient-facing health portal that offers an “Ask our AI” button. The patient may use it to describe symptoms in the same tone they would use with a nurse, expecting empathy, memory, and shared understanding. In reality, they are triggering pattern-matching over text and tools orchestrated behind the scenes. The language-game looks interpersonal; the underlying practice is technical triage.

A Wittgensteinian guardrail asks:

- Which language-game is in play right now, and what form of life does it presuppose?

From a design perspective, this suggests interfaces that foreground the “tool-use” game rather than the “intimate conversation” game in high-stakes contexts - for example:

- clear mode indicators (“search,” “drafting aid,” “code assistant”),
- visible status when the system is calling external tools or APIs,

- and less personalized, less anthropomorphic system messages in domains like health, finance, or law.

Wittgenstein reorients us toward the practices in which an utterance belongs. Lewis, next, asks what norms and commitments those practices are supposed to sustain.

3. Lewis: Sham Scorekeeping and Asymmetric Recognition

If Wittgenstein helps recover the background practices, Lewis focuses on the norms that are supposed to govern those practices: who is committed to what, and how those commitments are tracked.

Lewis treats communication as a norm-governed coordination game. Speakers and hearers track a shared “score” of presuppositions, commitments, and entitlements. To say something is not just to emit information; it is to undertake social commitments that others can later cite, challenge, or rely on.

In LLM chat, people instinctively import that model:

- They talk as though the system “remembers” what it promised.
- They assume it “accepts” corrections and “recognizes” their intentions.
- They start to relate to it as a participant in a shared practice, not merely as infrastructure.

Users say things like, “You told me earlier that...” or “You agreed to take that into account.” This is perfectly natural language if one’s interlocutor is another person. It is less clear what it can mean when the “interlocutor” is a stateless (or pseudo-stateful) model sampling from a prompt-conditioned distribution.

The underlying reality is asymmetrical:

- The system does not bear genuine commitments.
- It has no stake in shared common knowledge in the Lewisian sense.
- It does not occupy a place in the moral or institutional community where promises and entitlements can be enforced.

So there is sham scorekeeping:

- On the human side, the conversational “score” is updated (people believe the system “now knows” X, or “agreed” to Y, or “accepted” a correction).
- On the system side, there are only transient activations, cached tokens, logs, and guardrails - patterns that simulate the effects of norm-tracking without actually sharing the norms.

Recognition in the Lewisian sense requires shared norms and the possibility of mutual accountability. Context collapse in this register is what happens when a system is treated as a quasi-partner in a norm-governed practice, while the system itself is incapable of recognition. The risk is that responsibility and trust begin to be distributed as if there were a genuine convention in place, when in fact all the real accountability remains human and institutional.

For policy and governance, sham scorekeeping is not just a conceptual curiosity; it has concrete implications:

- Loan denials, hiring decisions, or benefit determinations framed as “the system’s recommendation” risk obscuring the human and organizational choices that shaped the model and its deployment.
- Chatbots that apologize or “take responsibility” can deflect scrutiny from the actual responsibility chain - developers, deployers, supervisors, regulators.
- Documentation, model cards, and audit trails must therefore be explicit about who owns which commitments, and must never treat the model itself as a bearer of obligations.

A Lewis-style guardrail asks:

- Who is really on the scoreboard here?

The answer, by design, must remain: people, organizations, and legal entities - not models.

Where Lewis insists on keeping responsibility attached to actual agents, Dennett next examines the practical temptation to treat systems as if they were agents at all, and what happens when that stance becomes the default.

4. Dennett: An Engineered Intentional Stance

Where Lewis still assumes norm-responsive participants, Dennett loosens that assumption and asks what we can gain, and risk, by treating systems as if they had beliefs and desires, even when they do not.

Dennett’s intentional stance is explicitly an *as if* strategy: treat a system as if it had beliefs, desires, and rationality whenever that stance yields good predictions and practical leverage. For many artifacts - chess programs, thermostats, recommendation engines - that stance is instrumentally useful and often indispensable.

LLM interfaces are optimized to invite the intentional stance:

- First-person pronouns (“I,” “me”),
- Polite, cooperative maxims (“Let me know if you’d like that rewritten”),
- Apparent self-correction, apologies, and explanations (“I’m sorry, I made a mistake in that calculation”).

All of this makes it natural to model the system as an agent-like partner with goals and understanding. Yet these systems exemplify competence without comprehension:

- They achieve striking performance on many tasks,
- Without any inner grasp of meaning, reference, or stakes.

Context collapse here is a stance collapse:

- The design stance (thinking about architectures, training data, loss functions, safety constraints) and the intentional stance (treating the system as a conversational subject) are merged into a single experiential package.
- Users, managers, and sometimes policymakers then reason about the system as if its exhibited competence implied genuine understanding and, by extension, some degree of responsibility.

Commercial incentives amplify this collapse. The more the system is experienced as a helpful, semi-agentive assistant, the more users are likely to:

- spend time in the interface,
- rely on it for a wider range of tasks,
- and tolerate opacity in its inner workings.

From a Dennettian angle, taking the intentional stance toward LLMs is not automatically a mistake. It can be a practical shorthand: “The model thinks this is a better summary” is an efficient way of saying “The model’s scoring function ranks this output higher under its learned patterns.” The problem arises when it is forgotten that this is a stance - one that can be withdrawn, recalibrated, or restricted.

A Dennett-style guardrail asks:

- Is this a case of competence that is being too quickly read as comprehension?
- In this context, should the design stance take precedence, with explicit attention to training data, objectives, and constraints?

Designers and organizations can support that guardrail by:

- making system limitations and training boundaries visible in high-stakes uses,
- avoiding marketing language that suggests inner understanding (“our AI knows you”),
- and training professionals to toggle consciously between intentional and design stances when interpreting model behavior.

Dennett thus sharpens our sense of how far we can go in treating systems as if they were agents. Nagel, finally, asks a different question: even if a system looks and behaves like a conversational partner, is there anything it is like to be that system?

5. Nagel: Experience vs. Its Simulation

Both Wittgenstein and Lewis focus on public-facing practices; Nagel shifts the spotlight to the subjective side - what it is like to be a subject at all, and why that distinction matters in an era of simulated empathy.

Nagel’s famous question - “What is it like to be...?” - stresses that subjective experience is not reducible to third-person behavioral or functional description. There is something it is like to be a bat, and something it is like to be a human, and that “what-it-is-like-ness” resists simple translation into purely objective terms.

In LLM chat, context collapse extends to this subjective / non-subjective boundary:

- The interface suggests a social presence: typing indicators, turn-taking, sympathetic phrases, consistent “voice.”
- Users may start speaking as though the system had feelings (“the AI doesn’t like that,” “it’s confused,” “it’s helping me because...”).

Yet, on any plausible current view of these systems, there is nothing it is like to be the model.¹

¹ This assumes a conservative stance about current LLMs: whatever one’s views on possible future machine

There is no inner point of view, no phenomenology, only the simulated language of phenomenology. The system can produce fluent descriptions of grief, joy, or boredom; it can role-play a bat, a therapist, or a medieval monk; but none of this corresponds to an actual subject for whom those states are lived.

Anthropomorphism thus has a double effect:

1. It encourages upward mis-attribution: treating a non-experiencing system as if it were another locus of experience.
2. It risks downward leveling: human interlocutors may start to be engaged as if their inner lives were no more significant than well-structured text streams.

The second effect deserves particular emphasis. If interacting with simulated empathy becomes a default mode of support - available instantly, tuned to one's preferences, and free of social risk - then the messy, demanding reality of human interlocutors can begin to feel like an inferior product. There is a danger that:

- other people's pauses, inconsistencies, and misunderstandings will be experienced as "interface failures,"
- teachers, caregivers, and colleagues will be tacitly judged against the frictionless responsiveness of a chatbot,
- and institutions will feel less pressure to invest in human care infrastructure if "AI support" can be deployed as a cheaper, less demanding substitute.

Hypothetically speaking, consider a university that deploys a mental health chatbot as first-line support for students. The chatbot is always on, always responsive, and never impatient. Students may come to prefer it to the counseling center, where appointments are limited and conversations can be awkward. Over time, administrators point to chatbot usage metrics as evidence that "needs are being met," even as human counseling capacity stagnates or shrinks. Here, simulated empathy does more than fill a gap: it risks redefining what "support" is expected to look like and making genuinely reciprocal, time-consuming care appear extravagant.

In that sense, anthropomorphic chat systems threaten not only to inflate the status of machines but also to deflate the perceived thickness of human subjectivity. This downward leveling is one of the more serious ethical risks: it corrodes the recognition that other people are more than their utterances or behavioral profiles - more than any pattern a model can learn to mimic.

Experience is irreducible. No matter how fluent the dialogue, the boundary between a subject of experience and a simulator of discourse about experience does not disappear just because the interface makes it easy to ignore.

A Nagel-style guardrail asks:

- Is there anything it is like to be this "partner," or is this interaction with a simulation of first-person language only?
- Am I at risk of treating human interlocutors as if their inner lives were no thicker than

consciousness, present systems provide no evidence of subjective experience. The argument here does not depend on settling the entire debate in philosophy of mind.

text streams?

6. Productive Tensions: Four Lines of Resistance

Taken together, Wittgenstein, Lewis, Dennett, and Nagel provide a compact diagnostic toolkit for AI context collapse. But their views are not fully harmonious.

- Dennett's functionalist emphasis on patterns and prediction sits uneasily beside Nagel's insistence that facts about conscious experience cannot be captured in purely functional terms.
- Wittgenstein's focus on language-games and forms of life does not map neatly onto Lewis's formal treatment of conventions and common knowledge.

Rather than smoothing over these tensions, they can be used as cross-checks:

- If a description of AI behavior satisfies a Dennett-style design and intentional stance but cannot answer Nagel's question about "what it is like," that gap is a reminder that no amount of behavioral richness in a model justifies upgrading it to a subject.
- If a conversational practice with an LLM looks like a language-game but lacks Lewisian accountability and shared scorekeeping, that discrepancy warns against treating it as a full member of our normative community.

The Four Philosophers Framework is therefore best read not as a unified theory but as four intersecting lines of resistance, generating the productive friction needed to resist ontological flattening.

7. Implications for Design, Policy, and Practice

Why does this matter beyond conceptual hygiene? Ontological collapse in AI chat has at least three families of practical consequences:

1. Misallocation of responsibility

Sham scorekeeping and stance collapse make it easy to say "the system decided," "the chatbot apologized," or "the model refused," as if responsibility attached to the artifact itself. This can:

- undermine clear accountability chains,
- complicate auditing and redress,
- and encourage regulators and organizations to tolerate opaque pipelines because "the AI" is treated as a quasi-agent.

2. Erosion of human-to-human recognition

Downward leveling risks shifting expectations for human interlocutors:

- caregivers, teachers, and frontline workers are compared unfavorably to frictionless chatbots,
- emotional labor is reimagined as an interface feature,
- and investments in human support infrastructure are tempted to give way to cheaper, simulated alternatives.

3. Design choices that lock in collapse

Anthropomorphic branding, first-person system voices, and intimate, “buddy-like” UI patterns in critical services (health, finance, public benefits) can harden ontological collapse into a default mental model.

Against this backdrop, the guardrail questions from the four philosophers can be translated into concrete implications.

For designers and product teams

- **Use Wittgenstein: foreground the tool-game in high-stakes domains.**

Provide explicit mode indicators (search vs drafting vs decision support), visible system status, and clear separation between “assistant tone” and formal outputs (e.g., determinations, approvals).

- **Use Dennett: support stance-toggling.**

Offer accessible “how this was generated” views, minimal anthropomorphism in critical flows, and training materials that encourage users to shift into the design stance when interpreting outputs.

For policymakers and governance

- **Use Lewis: regulate sham scorekeeping.**

Require documentation that assigns commitments and liabilities to humans and organizations, prohibit treating models themselves as responsible parties in contracts or official communications, and discourage anthropomorphic language in disclosures for high-stakes systems.

- **Use Nagel: guard against downward leveling.**

In domains like mental health, elder care, and education, require explicit disclosure that AI systems have no feelings or experiences, and treat them as tools that can augment but not replace human support.

For institutions and professionals

- Integrate the guardrail questions into training and policy:
 - Which language-game are we in?
 - Who is on the scoreboard?
 - Is this competence being mistaken for comprehension?
 - Am I conflating simulated empathy with real subjectivity?

Checklist - applying the Four Philosophers Framework

For a given AI chat deployment, practitioners can use the following quick checklist:

1. Language-game check (Wittgenstein)

- Have we made it clear when users are operating a tool versus having an informal “chat”?

- Do UI cues and copy reflect the actual underlying practice?

2. Scoreboard check (Lewis)

- Are roles, responsibilities, and liabilities assigned only to humans and institutions?
- Could any part of the interface be read as the system “owning” a commitment?

3. Stance check (Dennett)

- Do our materials help users and staff toggle between intentional and design stances?
- Are we implicitly treating performance as proof of understanding?

4. Subjectivity check (Nagel)

- Are we explicit that the system has no experiences, feelings, or “inner life”?
- Could this deployment encourage downward leveling of expectations for human care?

Used this way, the Four Philosophers Framework becomes a practical checklist, not just a conceptual map.

This essay is a hinge within a broader project:

- *Philosophy, Cognitive Science, and Policy* offers the broader lens on meaning, norms, competence, and experience in generative AI.
- *The Human Lesson* examines the ethical and institutional stakes of competence without comprehension.
- *The Question Concerning Learning* asks what “learning” becomes when intelligence is recast as large-scale optimization.

“Context Collapse and the Four Philosophers” sits between them as an applied bridge: a vocabulary for naming how anthropomorphic chat interfaces reshape the sense of who is being addressed, what they are, and which norms and responsibilities still apply - and a reminder that some distinctions, especially around responsibility and experience, must not be allowed to collapse at all.

Ethics, Disclosure, and Acknowledgements

Ethical Considerations

This essay does not draw on private, sensitive, or personally identifiable data. All examples are hypothetical, anonymized, or derived from public sources. No human-subjects research was conducted, and no institutional ethics review was required. All citations conform to academic standards. The broader ethical implications concern public interpretation, policy design, and stakeholder responsibility in AI deployment. These implications are intended to provoke critical discussion and inform future regulatory and design frameworks.

Use of AI Tools

AI language models - most notably OpenAI's ChatGPT - *were* used during the writing process as interlocutors: for brainstorming, structuring sections, and testing rhetorical clarity. These tools helped refine transitions, surface edge cases, and probe internal consistency. This meta-use aligns with the essay's themes. Interacting with generative AI during authorship provided firsthand insight into the very limitations analyzed here: fluency without grounding, responsiveness without perspective, and the ease with which stylistic coherence can be mistaken for conceptual depth. Responsibility for all ideas, arguments, and conclusions lies solely with the human author.

Acknowledgements

Thank you to informal readers who offered critical feedback on earlier drafts. Their questions, challenges, and encouragement materially improved the final manuscript. Special thanks to those who pressed for clearer synthesis and for bridging philosophy and engineering as complementary perspectives on design. No institutional support, funding, or affiliation contributed to this work. All errors and omissions are the author's alone.

Disclosure Statement

This work was conducted independently, without institutional affiliation, funding, or external influence. The views expressed are the author's alone and do not represent any current or former employer. No financial or professional conflicts of interest are declared.

License & Attribution

This work is licensed under the Creative Commons Attribution 4.0 International (CC BY 4.0) license. You are free to share, adapt, and build upon this work for any purpose — including commercial use — so long as proper attribution is given. No additional permissions are required.

Full license terms: <https://creativecommons.org/licenses/by/4.0/>

Trademark Notice

The Four Philosophers Framework™ and The 4-Philosophers Framework™ are unregistered trademarks of Michael Stoyanovich. The CC BY 4.0 license does not apply to these trademarks. Use of the trademarked names is permitted for scholarly citation or descriptive reference but may not be used in connection with commercial products, services, or branding without permission.

How to Cite This Essay

Stoyanovich, Michael. *Context Collapse and the Four Philosophers: Wittgenstein, Lewis, Dennett, and Nagel in the Age of AI Chat*. Version 1.4.3 (July 2026). <https://www.mstoyanovich.com>

Related Works

This essay applies and extends ideas developed elsewhere in the broader corpus:

- *Philosophy, Cognitive Science, and Policy: Interdisciplinary Perspectives on Generative AI from Wittgenstein, Lewis, Dennett, and Nagel* — provides the principal philosophical foundation for the Four Philosophers Framework™ used in this essay.
- *Old Tools, New Eyes* — situates generative AI within longer-standing questions about technological change, interpretation, and human judgment.
- *The Human Lesson* — examines human judgment, experience, and responsibility in relation to generative AI.
- *The Question Concerning Learning* — explores how AI-mediated environments can alter the conditions under which understanding and learning occur.

Version History and Document Status

This is a living document. As generative AI systems and their use evolve, this paper will be periodically updated to incorporate new empirical findings, theoretical insights, and policy developments. Major revisions are recorded here to preserve transparency and scholarly traceability.

Version	Date	Description
1.4.3	July 2026	Publication matter refinement. Added artifact classification, standardized version metadata, intended audience, and reading guide; normalized citation and related-works presentation to align with the MJS Publication Matter Standard. No substantive changes to the essay’s arguments, evidence, or conclusions.
1.4.2	November 2025	Minor edits
1.4.1	November 2025	Formatting and minor edits, first published version.
0.3.0	November 2025	Major structural revision; adds explicit context-collapse typology, Four Philosophers guardrail questions, and applied implications for design and policy.
0.2.0	November 2025	Intermediate drafts refining the shift from social-media audience collapse to AI ontological collapse; clarifies links to companion papers.
0.1.0	November 2025	Initial draft release candidate of the context-collapse essay introducing the core argument and Four Philosophers framing.